

УДК 004.85

БОЛЬШИЕ ЯЗЫКОВЫЕ МОДЕЛИ В ЗАДАЧЕ РАЗРЕШЕНИЯ ЛЕКСИЧЕСКОЙ МНОГОЗНАЧНОСТИ НА МАТЕРИАЛЕ РУССКОГО ЯЗЫКА

П. А. Гусяцкая¹, Н. В. Лукашевич²

¹⁻²*Московский государственный университет имени М. В. Ломоносова, г. Москва, Россия*

²*ИППИ РАН им. А. А. Харкевича*

¹*polinagousyatskaya@gmail.com, ²louk_nat@mail.ru*

Аннотация

Статья посвящена экспериментам в области автоматического разрешения неоднозначности (англ. *Word Sense Disambiguation, WSD*) на материале русского языка при помощи генеративных (decoder-only) моделей малого и среднего размера. Использование генеративных моделей напрямую не является оптимальным подходом для данной задачи, однако такие модели имеют потенциал в роли семантических разметчиков необработанных данных. Автоматизация семантической разметки текстов при помощи генеративных моделей потенциально способна преодолеть бутылочное горлышко недостатка размеченных данных для обучения энкодеров.

Как показывают более ранние исследования, флагманские англоязычные и мультиязычные модели способны достичь более 90% аккуратности на данной задаче, модели меньшего размера – 80%+. Настоящее исследование ставит своей целью установить, решается ли аналогичная задача на материале русского языка с помощью русифицированных моделей малого и среднего размера (до 32B), не требующих большого количества вычислительных ресурсов для использования.

Эксперимент по разрешению неоднозначности в данной работе проводится как в базовой постановке (one/few-shot prompting), так и в различных модификациях (обогащение контекста словарной информацией – гиперонимами, гипонимами, метками тематической области и т.д., анализ широкого и узкого контекстного окна неоднозначной лексемы, ансамблевые

подходы, в которых одна модель валидирует и корректирует предсказания другой). В качестве материала исследования используется русскоязычный размеченный ресурс RuSemCor, семантическая разметка которого соответствует категориям семантической сети RuWordNet.

По результатам экспериментов модели показывают себя пригодными для задачи: все модели выходят за уровень случайного предсказания, а наиболее мощные достигают 80% аккуратности, сопоставимых с результатами англоязычных моделей того же размера. Более информативным для моделей показывает себя широкий контекст неоднозначной лексики. Подходы с дообогащением входных данных и ансамблевые методы дают значительный прирост в качестве.

Ключевые слова: *Разрешение лексической неоднозначности, компьютерная семантика, большие языковые модели.*

ВВЕДЕНИЕ

Целью данной работы является оценка эффективности русифицированных генеративных моделей малого и среднего размера на задаче автоматического разрешения неоднозначности (англ. *Word Sense Disambiguation, WSD*). Отметим, что генеративные модели не являются оптимальными для классификационных задач: они подвержены галлюцинациям, испытывают трудности с выводом структурного ответа, а также проигрывают по времени работы. Модели такого типа, однако, обладают двумя весомыми преимуществами. Во-первых, они способны решать различные задачи без предварительного дообучения; во-вторых, они обладают способностью к рассуждению, не реализованной в качестве отдельного модуля логики, а являющейся следствием кодирования информации и ее трансформации на слоях внимания. Все это делает генеративные модели перспективными в роли семантических разметчиков больших объемов необработанного текста, потенциально способных преодолеть бутылочное горлышко недостатка размеченных данных для машинного обучения.

Глобальная цель экспериментов в этой области – разработка инструмента для семантической разметки русскоязычных текстов, качество разметки которого будет сопоставимо с качеством аннотатора-человека.

АНАЛОГИЧНЫЕ ИССЛЕДОВАНИЯ

В современной компьютерной лингвистике проблема автоматического разрешения неоднозначности с использованием генеративных моделей находится на начальном этапе проработки. Во многом, исследования в этой области зависят от доступности корпусов с семантической разметкой, в особенности – отражающей систему значений многозначных слов. Для английского языка существуют такие ресурсы, как WordNet [Miller 1995], BabelNet [Navigli, Ponzetto 2010] или SyntagNet [Maru et al. 2019] – их основная особенность заключается в том, что представление каждого из значений многозначной леммы содержит множество синтагматических и парадигматических связей, в которые оно вступает. Данное исследование основывается на аналогичном русскоязычном ресурсе – RuSemCor, семантически размеченном в рамках инвентаря значений русскоязычной семантической сети RuWordNet [Loukachevitch et al. 2016].

Как упоминалось ранее, применение генеративных моделей к задаче разрешения лексической многозначности напрямую не является оптимальным подходом. Генеративные модели, однако, рассматриваются как потенциальный источник радикальных, или прорывных инноваций (англ. *disruptive innovation*) – а именно, технологий, которые могут обеспечить резкий прорыв в исследованиях в какой-либо области, существенно превосходя ранее известные методы и подходы по качеству [Yae et al 2025]. Именно это послужило толчком для применения моделей такого типа в том числе и к задачам автоматической обработки языка.

Известно несколько фундаментальных исследований генеративных моделей в задаче разрешения лексико-семантической неоднозначности. В работе [там же] тестируется зависимость качества определения нужного значения полисеманта от различных техник промптинга: модели предлагается или выбрать значение слова из предложенного списка, или оценить истинность суждения типа «в предложении X слово Y употреблено в значении Z». В первом подходе крупные модели (GPT 4) достигают 70% по метрике аккуратности, более легкие модели (Llama-2 7B, 13B) останавливаются на пороге качества случайного предсказания (42.56%, 50.5%). Во втором подходе модели справляются немного лучше – от 53.5% на модели в 7 млрд. параметров до 77% на флагманской модели. Помимо всего прочего, в данном исследовании формулируется фундаментальная закономерность: качество классификации находится в прямой зависимости от количества параметров модели. Исследование [Sumanathilaka et al. 2025] ставит своей целью подобрать оптимальную стратегию промптинга для задачи разрешения неоднозначности: тестировались техники промпт-инжиниринга, zero-shot и few-shot подходы, а также режим эксплицитного последовательного рассуждения, получивший название Chain-of-Thought. Последний подход показал себя как самый результативный для задачи; для дополнительного улучшения качества авторы прибегли к обогащению описаний значений полисемантов синонимами. Вкупе с флагманской моделью GPT-4-Turbo подход показал рекордные 90% по метрике аккуратности. Базовый уровень для исследований на русскоязычном материале намечается в работе [Kirillovich et al. 2025]; Эксперименты по разрешению многозначности проводятся на ряде крупных языковых моделей – как в базовой постановке, так и с применением техник обогащения входных данных словарной информацией. Флагманские

англоязычные и мультязычные модели (GPT-4, Deepseek) достигают качества в 90% и более по метрике аккуратности, модели меньшего размера достигают 80+. Исследования в области разрешения неоднозначности с использованием флагманских генеративных моделей внесли большой вклад в развитие этого направления – однако, учитывая дороговизну и ресурсозатратность флагманских моделей, такие решения нельзя признать целевыми. Одним из основных вызовов данной сферы, таким образом, остается достижение конкурентноспособного качества на моделях меньшего размера.

Эксперименты по проверке информативности широкого и узкого контекста целевого слова в задаче WSD проводились как в ранних подходах, основанных на знаниях [Agirre, Edmonds 2007], так и в нейросетевых подходах (к примеру [Митрофанова и др. 2013]). Основной исследовательский интерес представляет вопрос, какой контекст содержит больше информативных контекстных маркеров, необходимых для разрешения неоднозначности – глобальный (документ, тематическая область) или локальный (словосочетания). Исследования контекстного окна в сценарии с генеративными моделями имеет значение с точки зрения экономии токенов: если высокая точность достигается на узких контекстах, привлечение широкого оказывается неоправданным, так как будет увеличивать стоимость и время работы модели за счёт увеличения промпта без значимого прироста к качеству.

НАБОР ДАННЫХ

Материалом исследования послужил корпус русскоязычных текстов RuSemCor. Данный ресурс представляет собой коллекцию текстов, семантически размеченных согласно инвентарю значений русскоязычной семантической сети RuWordNet. Структура ресурса такова: каждой лексеме соответствует одно или несколько значений, определенных по идентификатору синсета и его краткому описанию; каждому значению соответствует список предложений, в которых реализуется именно это значение.

Эксперименты проводились на части RuSemCor: из корпуса были отобраны неоднозначные лексемы, каждое из значений которых было представлено хотя бы одним предложением. Предложения из корпуса были сгруппированы по неоднозначным лексемам и их значениям – пример

структуры данных представлен на Рис. 1. Также, ввиду того, что эксперименты с LLM занимают довольно долгое время, для первоначальных тестов набор данных был урезан до 3000 предложений.

```
"акцент": {
  "107128-N-132865": {
    "RwnSynsetId": "107128-N",
    "RwnSynsetName": "АКЦЕНТ (ПРОИЗНОШЕНИЕ)",
    "sentences": [
      {
        "SentenceId": 66066,
        "Sentence": "Ладно бы еще переозвучивали тех, кто говорит с акцентом - Баниониса, Брыльску, Ярвета."
      }
    ]
  },
  "115881-N-132865": {
    "RwnSynsetId": "115881-N",
    "RwnSynsetName": "ПОДЧЕРКНУТЬ, СДЕЛАТЬ АКЦЕНТ",
    "sentences": [
      {
        "SentenceId": 46418,
        "Sentence": "Печально, но рг [в довольно широком смысле] – это та деятельность, которой необходимо заниматься регулярно и систематически, причём с акцентом на слово необходимо."
      }
    ]
  }
}
```

Рис. 1. Пример структуры данных из RuSemCor, использующейся в экспериментах

МОДЕЛИ

Для экспериментов было отобрано две группы генеративных моделей. В первую группу вошли две модели, не превышающие 12 миллиардов параметров – saiga-llama3 (8b) и Vikhr-Nemo (12B). В группу моделей среднего размера вошли три модели, насчитывающие от 24 до 32 миллиардов параметров – модель семейства Vikhr Vistral-24B-Instruct, Mistral-Small-24B-Instruct и RuadaptQwen3-32B-Instruct. Все модели доступны на ресурсе HuggingFace через библиотеку transformers.

В экспериментах использовались только модели типа Instruct, специально обученные на точное выполнение инструкций пользователя. Выбор моделей такого типа обусловлен необходимостью получать от моделей структурный ответ в формате json, чтобы иметь возможность автоматизировать оценку качества классификации. Модели типа Base, работающие по принципу генерации наиболее вероятного следующего

токена в цепочке, для классификационных задач, масштабированных на объемные наборы данных, не подходят.

ОПИСАНИЕ ЭКСПЕРИМЕНТОВ

В рамках проведенных экспериментов моделям на вход предлагалось описание задачи и один пример ее решения (англ. one-shot prompting). Кроме базовой постановки задачи были протестированы некоторые подходы к улучшению качества классификации – в частности, дообогащение входных данных синтагматической и парадигматической информацией, а также ансамблевый подход, где одна модель валидирует и исправляет ответы другой.

В базовой постановке задачи на вход модели подавалось предложение, содержащее неоднозначную лексему и список ее возможных значений с индексом значения и кратким описанием, соответствующим наименованию соответствующего синсета в RuWordNet. Пример структуры данных, подаваемой на вход, представлен на Рис. 2.

```
{
  "lexeme": "продукт",
  "sentence_id": 47901,
  "sentence_text": "Нарушение правил хранения, закупки или рационального использования зерна и продуктов его переработки, а также правил производства продуктов переработки зерна",
  "senses": [
    "106505-N-149796 ПРОДУКТ ТРУДА",
    "106507-N-149796 ПРОДУКТ ПРОИЗВОДСТВА",
    "368-N-149796 ПРОДУКТЫ ПИТАНИЯ",
    "106482-N-149796 РЕЗУЛЬТАТ, ИТОГ"
  ]
}
```

Рис. 2. Входные данные при базовой постановке задачи (полный контекст).

Базовый эксперимент проводился с использованием полного контекста неоднозначной лексики (в среднем 19.5 токенов), а также с использованием узкого контекста, когда входное предложение обрезалось до контекстного окна [-2; +2] относительно целевой лексики, как видно на примере:

(a) "sentence_text": "Пользователи Facebook в скором <TERM>времени</TERM> смогут общаться по Skype через Facebook Connect."

(b) "windowed_context": "в скором <TERM>времени</TERM> смогут общаться"

В подходах с обогащением входных данных краткое описание значения дополнялось списками гипонимов, гиперонимов и меток домена, извлеченных из ресурса RuWordNet для каждого значения целевой лексики

(примеры данных представлены на Рис. 3). Модели предлагалось принять решение, обращая внимание на дополнительную словарную информацию. Стоит отметить, что в данном эксперименте в контекст модели привлекается не вся словарная информация, представленная в RuWordNet, а только категории, имеющиеся у большинства лексем. Некоторые категории, такие как синонимы в другой части речи, антонимы, причины, следствия, меронимы или холонимы, у большинства лексем, составляющих RuSemCor, являются пустыми, в связи с чем было принято решение не включать их в эксперимент.

```

"domains": {
  "113869-N-154364": [
    "ЛИТЕРАТУРА",
    "ФИЛОЛОГИЯ"
  ],
  "122006-N-154364": [
    "ИНТЕРНЕТ",
    "КОМПЬЮТЕР"
  ]
},

"hyponyms": {
  "113869-N-154364": [
    "ПОЛОСА ГАЗЕТЫ",
    "РАЗВОРОТ ИЗДАНИЯ, ТЕТРАДИ",
    "ФОРЗАЦ КНИГИ"
  ],
  "122006-N-154364": [
    "ДОМАШНЯЯ СТРАНИЦА"
  ]
},

"hypernyms": {
  "113869-N-154364": [
    "ЛИСТ БУМАГИ"
  ],
  "122006-N-154364": [
    "ИЗОБРАЖЕНИЕ (РЕЗУЛЬТАТ)"
  ]
}

```

Рис. 3. Пример обогащения лексемы «страница» метками домена (тематической области), гипонимами и гиперонимами

Ансамблевый подход предполагал следующий сценарий: модель-разметчик проходила по набору данных, выводя в ответ метку предложения, метку нужного значения, а также краткое рассуждение, почему она приняла именно такое решение. Модели-валидатору предлагалось оценить ответы и рассуждения модели-разметчика и исправить их, если решение было оценено как неверное.

```
#разметчик
{
  "lexeme": "брать",
  "sentence_id": 108619,
  "sense_id": "115594-V-133880",
  "reason": "Целевое слово используется в контексте ручного действия, что соответствует значению \"брать
рукой\""
}
#валидатор
{
  "score": 1,
  "reeval_sense_id": null,
  "reason": "Контекст предложения указывает на использование ложки для взятия теста, что соответствует
значению БРАТЬ РУКОЙ."
```

Рис. 4. Пример взаимодействия модели-разметчика и модели-валидатора в ансамблевом подходе.

РЕЗУЛЬТАТЫ

Качество классификации оценивалось по метрике ассурасу. Результаты моделей по базовому сценарию и сценарию с дообогащением представлены в Табл. 1. В Табл. 2 представлены результаты ансамблевых подходов.

Таблица 1. Результаты классификации

	saiga-llama 8B	vikhr-nemo 12B	mistral 24B	vikhr 24B	ruadapt 32B
Базовый сцен. + широкий конт.	0.58	0.64	0.74	0.73	0.76
Базовый сцен. + узкий конт.	0.55	0.61	0.73	0.71	0.74
Дообогащение	0.59	0.65	0.779	0.806	0.8

Таблица 2. Результаты ансамблевых подходов к классификации

	одиночная модель	ансамблевые подходы	
	saiga-llama 8B	vikhr-nemo12B + saiga_llama 8B	mistral 24B + saiga-llama 8B
	acc	acc	acc
Базовый сцен.	0.58	0.64	0.74

Полученные результаты позволяют сделать следующие выводы:

1) Задача автоматического разрешения неоднозначности русскоязычных существительных в целом представляется посильной для русифицированных или мультязычных моделей больше 8 миллиардов параметров. Даже модели в 8 млрд. показывают качество, превышающее

качество случайного выбора, а модели в 24 и больше млрд. достигают качества 80%+, приближающегося к отметке, установленной на крупных англоязычных моделях [Kirillovich et al. 2025].

2) Обоогащение входных данных дополнительной словарной информацией дает прирост к качеству классификации, в особенности на более крупных моделях, что позволяет признать такой подход целевым для задачи WSD с использованием больших языковых моделей.

3) Короткий контекст показывает себя менее информативным для генеративных моделей: по всей видимости, ближайшее окружение неоднозначной лексики не содержит достаточного количества контекстных маркеров для успешного разрешения неоднозначности. Глубинный анализ причин таких результатов, однако, требует дальнейшего исследования.

4) Ансамблевые подходы, где модель-валидатор исправляет ответы модели-первичного разметчика, также показывают заметный прирост к качеству. Перспективным сценарием – как с точки зрения ресурсов, так и с точки зрения качества – является использование более крупной модели в качестве валидатора более легкой. На наших данных ансамблевые методы дают до 16% прироста по метрикам.

ЗАКЛЮЧЕНИЕ

В рамках настоящего исследования были проведены эксперименты по автоматическому разрешению неоднозначности русскоязычных лексем с помощью генеративных моделей малого и среднего размера. Результаты экспериментов показывают, что средние русифицированные модели способны осуществлять автоматическую классификацию с уровнем качества 80+, что приближается к целевому уровню, заданному флагманскими моделями. Положительную динамику показывают ансамблевые методы классификации, а также подходы с обогащением контекста модели лингвистической информацией. Полученные результаты закладывают основу для разработки масштабного инструмента для автоматической семантической разметки больших объемов текстов.

Благодарности

Работа выполнена при поддержке гранта РФФИ №24-18-00988.

СПИСОК ЛИТЕРАТУРЫ

1. Митрофанова О. А. и др. АВТОМАТИЧЕСКОЕ РАЗРЕШЕНИЕ ЛЕКСИКО-СЕМАНТИЧЕСКОЙ НЕОДНОЗНАЧНОСТИ И ВЫДЕЛЕНИЕ КОНСТРУКЦИЙ (на материале Национального корпуса русского языка) //Лексикология. Лексикография и Корпусная лингвистика. – 2013. – С. 122-143.
2. Agirre E., Edmonds P. (ed.). Word sense disambiguation: Algorithms and applications. – Springer Science & Business Media, 2007. – Т. 33.
3. Kirillovich A. et al. RuSemCor: A Word Sense Disambiguation corpus for Russian //Proceedings of the 34th ACM International Conference on Information and Knowledge Management. – 2025. – С. 6438-6443.
4. Loukachevitch N. V., Lashevich G., Gerasimova A. A., Ivanov V. V., Dobrov B. V. Creating Russian WordNet by Conversion // In Proceedings of Conference on Computational linguistics and Intellectual technologies Dialog-2016, 2016. pp. 405-415
5. Maru M. et al. SyntagNet: Challenging supervised word sense disambiguation with lexical-semantic combinations //Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP). – 2019. – С. 3534-3540.
6. Miller G. WordNet: A Lexical Database for English. Communications of the ACM - Vol. 38, No. 11 - 1995 - pp. 39-41.
7. Navigli R., Ponzetto S. P. BabelNet: Building a very large multilingual semantic network //Proceedings of the 48th annual meeting of the association for computational linguistics. – 2010. – С. 216-225.
8. Sumanathilaka D., Micallef N., Hough J. Glossgpt: Gpt for word sense disambiguation using few-shot chain-of-thought prompting //Procedia Computer Science. – 2025. – Т. 257. – С. 785-792.
9. Yae J. H. et al. Leveraging large language models for word sense disambiguation //Neural Computing and Applications. – 2025. – Т. 37. – №. 6. – С. 4093-4110.

МОДЕЛИ

1. IlyaGusev/saiga_llama3_8b
URL: https://huggingface.co/IlyaGusev/saiga_llama3_8b (дата обращения 29.12.2025)
2. mistralai/Mistral-Small-24B-Instruct-2501
URL: <https://huggingface.co/mistralai/Mistral-Small-24B-Instruct-2501> (дата обращения 29.12.2025)
3. RefalMachine/RuadaptQwen3-32B-Instruct
URL: <https://huggingface.co/RefalMachine/RuadaptQwen3-32B-Instruct> (дата обращения 29.12.2025)
4. Vikhrmodels/Vikhr-Nemo-12B-Instruct-R-21-09-24
URL: <https://huggingface.co/Vikhrmodels/Vikhr-Nemo-12B-Instruct-R-21-09-24> (дата обращения 29.12.2025)
5. Vikhrmodels/Vistral-24B-Instruct
URL: <https://huggingface.co/Vikhrmodels/Vistral-24B-Instruct> (дата обращения 29.12.2025)

LLM APPLICATIONS TO THE TASK OF WORD SENSE DISAMBIGUATION IN RUSSIAN TEXTS

P. A. Gousyatskaya¹, N. V. Loukachevitch²

Lomonosov Moscow State University, Moscow, Russia

¹polinagousyatskaya@gmail.com, ²louk_nat@mail.ru

Abstract

This study is dedicated to experiments in Word Sense Disambiguation on Russian-language material using generative (decoder-only) models of small and medium size. While the direct application of generative models is not an optimal approach for this task, such models demonstrate potential as semantic annotators of large volumes of raw data. The automation of semantic text annotation using generative models may help overcome the bottleneck of insufficient labeled data for training encoder-based models.

Previous studies show that state-of-the-art English and multilingual models can achieve accuracy above 90% on this task, while smaller models typically reach 80%+. The present study aims to determine whether a comparable task can be successfully addressed for small- and medium-scale models adapted to Russian, that do not require substantial computational resources.

The experiments reported in this paper are conducted both in a basic setting (one/few-shot prompting) with separate tests for narrow and broad context window of the ambiguous lexeme, as well as under several modifications, such as context enrichment with paradigmatic information (e.g., hypernyms, hyponyms, domain labels of the target word) and ensemble approaches in which one model validates and refines the predictions of another. The study is based on the Russian annotated resource RuSemCor, annotated in terms of the RuWordNet semantic net.

The experiments ultimately pursue the development of an efficient and accessible pipeline for automatic semantic annotation of Russian texts, useful both for direct application and for preparing annotated data for machine-learning tasks.

Keywords: *Word Sense Disambiguation, LLM, lexical ambiguity, classification, computational semantics*

REFERENCES

1. Mitrofanova O. A. et al. Avtomaticheskoe razreshenie leksiko-semantichekoi neodnoznachnosti i vydelenie konstruktsii (na materiale Natsional'nogo korpusa russkogo yazyka) [Automatic word sense disambiguation and construction extraction (based on the Russian National Corpus)] // Leksikologiya. Leksikografiya i Korpusnaya Lingvistika [Lexicology. Lexicography and Corpus Linguistics]. 2013. P. 122–143
2. Kirillovich A. et al. RuSemCor: A Word Sense Disambiguation corpus for Russian // Proceedings of the 34th ACM International Conference on Information and Knowledge Management. 2025. P. 6438-6443.
3. Loukachevitch N. V., Lashevich G., Gerasimova A. A., Ivanov V. V., Dobrov B. V. Creating Russian WordNet by Conversion // In Proceedings of Conference on Computational linguistics and Intellectual technologies Dialog-2016, 2016. pp. 405-415
4. Maru M. et al. SyntagNet: Challenging supervised word sense disambiguation with lexical-semantic combinations // Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP). – 2019. – C. 3534-3540.
5. Miller G. WordNet: A Lexical Database for English. Communications of the ACM - Vol. 38, No. 11. 1995. P. 39-41.
6. Navigli R., Ponzetto S. P. BabelNet: Building a very large multilingual semantic network // Proceedings of the 48th annual meeting of the association for computational linguistics. – 2010. – C. 216-225.
7. Sumanathilaka D., Micallef N., Hough J. Glossgpt: Gpt for word sense disambiguation using few-shot chain-of-thought prompting // Procedia Computer Science. – 2025. – T. 257. – C. 785-792.
8. Yae J. H. et al. Leveraging large language models for word sense disambiguation // Neural Computing and Applications. – 2025. – T. 37. – №. 6. – C. 4093-4110.

MODELS

1. IlyaGusev/saiga_llama3_8b
URL: https://huggingface.co/IlyaGusev/saiga_llama3_8b (accessed 29.12.2025)
2. mistralai/Mistral-Small-24B-Instruct-2501
URL: <https://huggingface.co/mistralai/Mistral-Small-24B-Instruct-2501> (accessed 29.12.2025)

3. RefalMachine/RuadaptQwen3-32B-Instruct

URL: <https://huggingface.co/RefalMachine/RuadaptQwen3-32B-Instruct>
(accessed 29.12.2025)

4. Vikhrmodels/Vikhr-Nemo-12B-Instruct-R-21-09-24

URL: <https://huggingface.co/Vikhrmodels/Vikhr-Nemo-12B-Instruct-R-21-09-24>
(accessed 29.12.2025)

5. Vikhrmodels/Vistral-24B-Instruct

URL: <https://huggingface.co/Vikhrmodels/Vistral-24B-Instruct> (accessed 29.12.2025)

СВЕДЕНИЯ ОБ АВТОРАХ

Полина Андреевна Гусяцкая – аспирант кафедры теоретической и прикладной лингвистики филологического факультета МГУ им. М. В. Ломоносова

Polina A. Gousyatskaya – PhD Student, Department of Theoretical and Applied Linguistics, Faculty of Philology, Lomonosov Moscow State University

Наталья Валентиновна Лукашевич – д.т.н., профессор кафедры теоретической и прикладной лингвистики филологического факультета МГУ им. М. В. Ломоносова

Natalia V. Loukachevitch – Doctor of Sciences, professor, Department of Theoretical and Applied Linguistics, Faculty of Philology, Lomonosov Moscow State University